

# Longitudinal Linear Models - Some Theory

## Terminology

- 1) Subjects - Clusters
- 2) Occasions - individual observations within clusters.
- 3) Levels of variables:
  - Time varying - micro - level 1 - vary by occasion  
e.g. Age
  - Time invariant - macro - level 2 - vary by subject  
- "contextual" variables  
e.g. Gender

## Simple example

- Time-varying variables
- All variables have random effects.

J subjects measured on T occasions.

Model for jth subject:

(Level I Model)

$$\underset{n_j \times 1}{\tilde{Y}_j} = \underset{n_j \times p}{X_j} \underset{p \times 1}{\tilde{\beta}_j} + \underset{n_j \times 1}{\tilde{\epsilon}_j}$$

$$\tilde{\epsilon}_j \sim N(\underline{0}, \sigma^2 I)$$

BLUE of  $\tilde{\beta}_j$  using  $\tilde{Y}_j$  and  $X_j$

$$\hat{\tilde{\beta}}_j = (X_j' X_j)^{-1} X_j' \tilde{Y}_j$$

Best among Linear Unbiased  
Estimators  
repeatedly sampling  
from subject j

$$E(\hat{\beta}_j | \beta_j) = \beta_j$$

$$\text{Var}(\hat{\beta}_j | \beta_j) = \sigma^2 (X_j' X_j)^{-1}$$


---

Model for  $\beta_j$  (Level 2 Model)

$$\beta_j \sim N(\gamma, G)$$

or

$$\beta_j = \gamma + \delta_j$$

$$\delta_j \sim N(0, G)$$

So

$$y_j = X_j \beta_j + \varepsilon_j = X_j (\gamma + \delta_j) + \varepsilon_j$$

$$= \underbrace{X_j \gamma}_{\text{fixed effects}} + \underbrace{X_j \delta_j + \varepsilon_j}_{\text{random effects}}$$

Putting clusters together:

$$\begin{bmatrix} y_1 \\ \tilde{y}_1 \\ y_2 \\ \tilde{y}_2 \\ \vdots \\ y_j \\ \tilde{y}_j \\ \vdots \\ y_J \\ \tilde{y}_J \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_j \\ \vdots \\ x_J \end{bmatrix} \gamma + ?$$

Putting clusters together:

$$\begin{bmatrix} y_1 \\ \tilde{y}_1 \\ y_2 \\ \tilde{y}_2 \\ \vdots \\ y_j \\ \tilde{y}_j \\ \vdots \\ y_J \\ \tilde{y}_J \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_j \\ \vdots \\ x_J \end{bmatrix} \gamma + \begin{bmatrix} x_1 \delta_1 \\ x_2 \delta_2 \\ \vdots \\ x_j \delta_j \\ \vdots \\ x_J \delta_J \end{bmatrix} + \begin{bmatrix} \tilde{y}_1 \\ \tilde{y}_2 \\ \vdots \\ \tilde{y}_j \\ \vdots \\ \tilde{y}_J \end{bmatrix}$$

Putting clusters together:

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_j \\ \vdots \\ y_J \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_j \\ \vdots \\ x_J \end{bmatrix} + \begin{bmatrix} z_1 & 0 & 0 & \dots & 0 \\ 0 & z_2 & 0 & & 0 \\ 0 & 0 & \ddots & & \\ \vdots & & & z_j & \\ 0 & & & & \ddots \\ & & & & & z_J \end{bmatrix} \begin{bmatrix} \delta_1 \\ \delta_2 \\ \vdots \\ \delta_j \\ \vdots \\ \delta_J \end{bmatrix} + \begin{bmatrix} \eta_1 \\ \eta_2 \\ \vdots \\ \eta_j \\ \vdots \\ \eta_J \end{bmatrix}$$

Putting clusters together:

$$\begin{bmatrix} y_1 \\ \tilde{y}_1 \\ y_2 \\ \tilde{y}_2 \\ \vdots \\ y_j \\ \tilde{y}_j \\ \vdots \\ y_J \\ \tilde{y}_J \end{bmatrix} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_j \\ \vdots \\ x_J \end{bmatrix} \gamma + \begin{bmatrix} z_1 & 0 & 0 & \dots & 0 \\ 0 & z_2 & 0 & & 0 \\ \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & 0 & & z_j & \\ \vdots & \vdots & & \vdots & \ddots \\ 0 & & & & z_J \end{bmatrix} \begin{bmatrix} \delta_1 \\ \tilde{\delta}_1 \\ \delta_2 \\ \tilde{\delta}_2 \\ \vdots \\ \delta_j \\ \tilde{\delta}_j \\ \vdots \\ \delta_J \\ \tilde{\delta}_J \end{bmatrix} + \begin{bmatrix} \varepsilon_1 \\ \tilde{\varepsilon}_1 \\ \varepsilon_2 \\ \tilde{\varepsilon}_2 \\ \vdots \\ \varepsilon_j \\ \tilde{\varepsilon}_j \\ \vdots \\ \varepsilon_J \\ \tilde{\varepsilon}_J \end{bmatrix}$$

$$\tilde{y} = X \gamma + Z \tilde{\delta} + \tilde{\varepsilon}$$

We use  $Z_i$  instead of  $X_i$  to anticipate that  $Z_i$  may differ from  $X_i$ .

Often  $T$  is large and  $\gamma$ ,  $G$ ,  $\sigma^2$  are estimated with much more precision than  $\hat{\beta}_j$ .

Pretend for this discussion that  $\gamma$ ,  $G$ ,  $\sigma^2$  are "known".

Can we improve on  $\hat{\beta}_j$  as an estimator of  $\beta_j$ ?

What do we know about  $\beta_j$ ?

1)  $E(\hat{\beta}_j - \beta_j) = 0$  with  $\text{Var}(\hat{\beta}_j - \beta_j) = \sigma^2 (X_j' X_j)^{-1}$

2)  $\beta_j$  comes from a population with mean  $\gamma$  and variance  $G$ .

$$\text{So } E(\underset{\sim}{y} - \underset{\sim}{\beta}_j) = \underset{\sim}{0} \text{ with } \text{Var}(\underset{\sim}{y} - \underset{\sim}{\beta}_j) = G$$

So  $\underset{\sim}{\hat{\beta}}_j$  &  $\underset{\sim}{y}$  are unbiased predictors of  $\underset{\sim}{\beta}_j$   
(treating  $\underset{\sim}{\beta}_j$  as random)

How to combine them?

Best Linear Unbiased Linear Combination Predictor

BLUP: Take a weighted average  
with weights proportional to "precision" =  $(\text{variance})^{-1}$

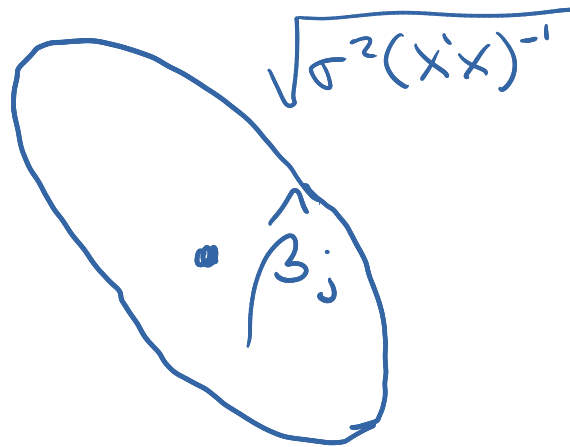
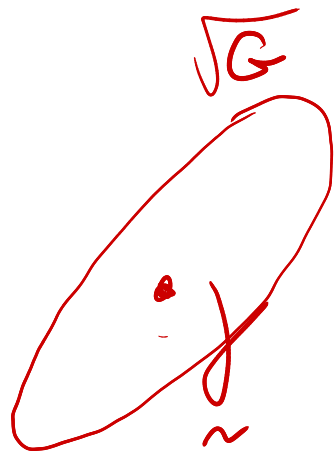
BLUP

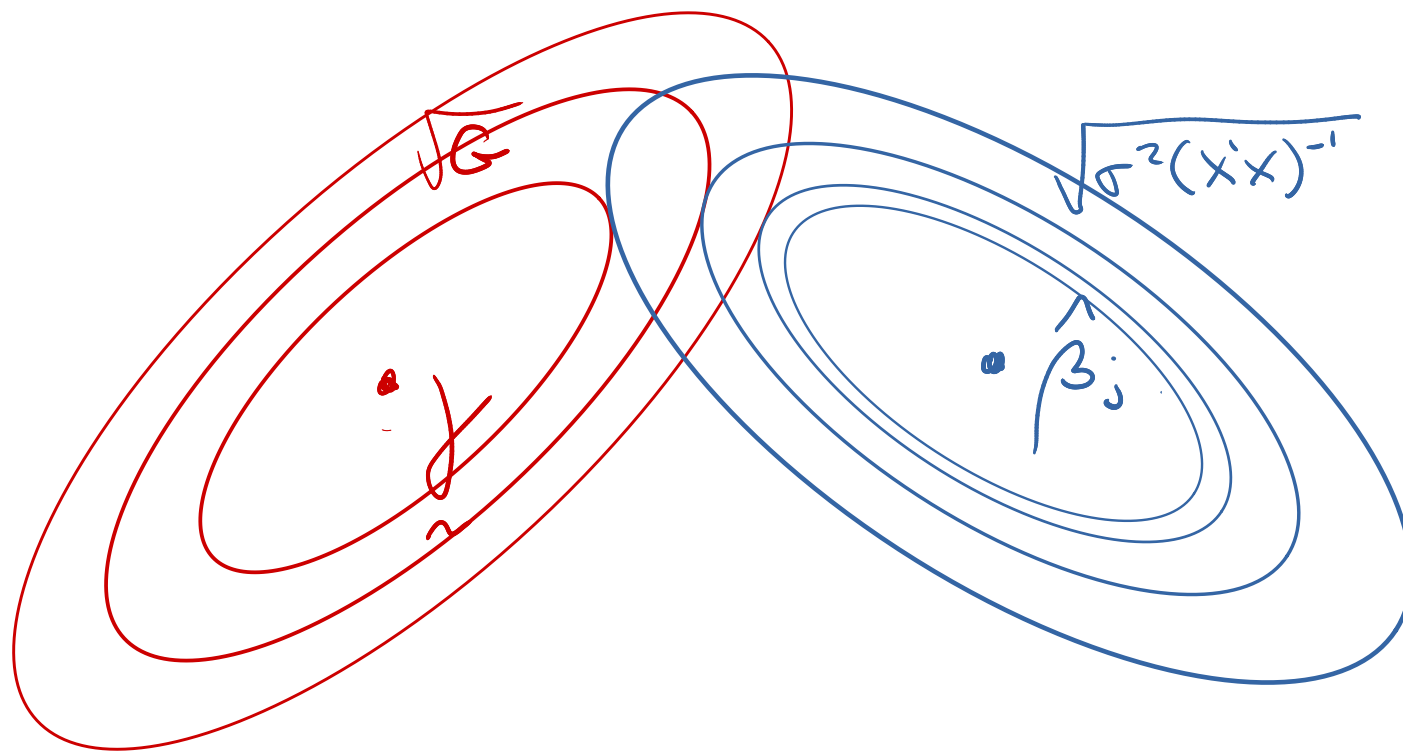
$$\tilde{\beta}_j = (W_1 + W_2)^{-1} (W_1 \tilde{y} + W_2 \hat{\beta}_j)$$

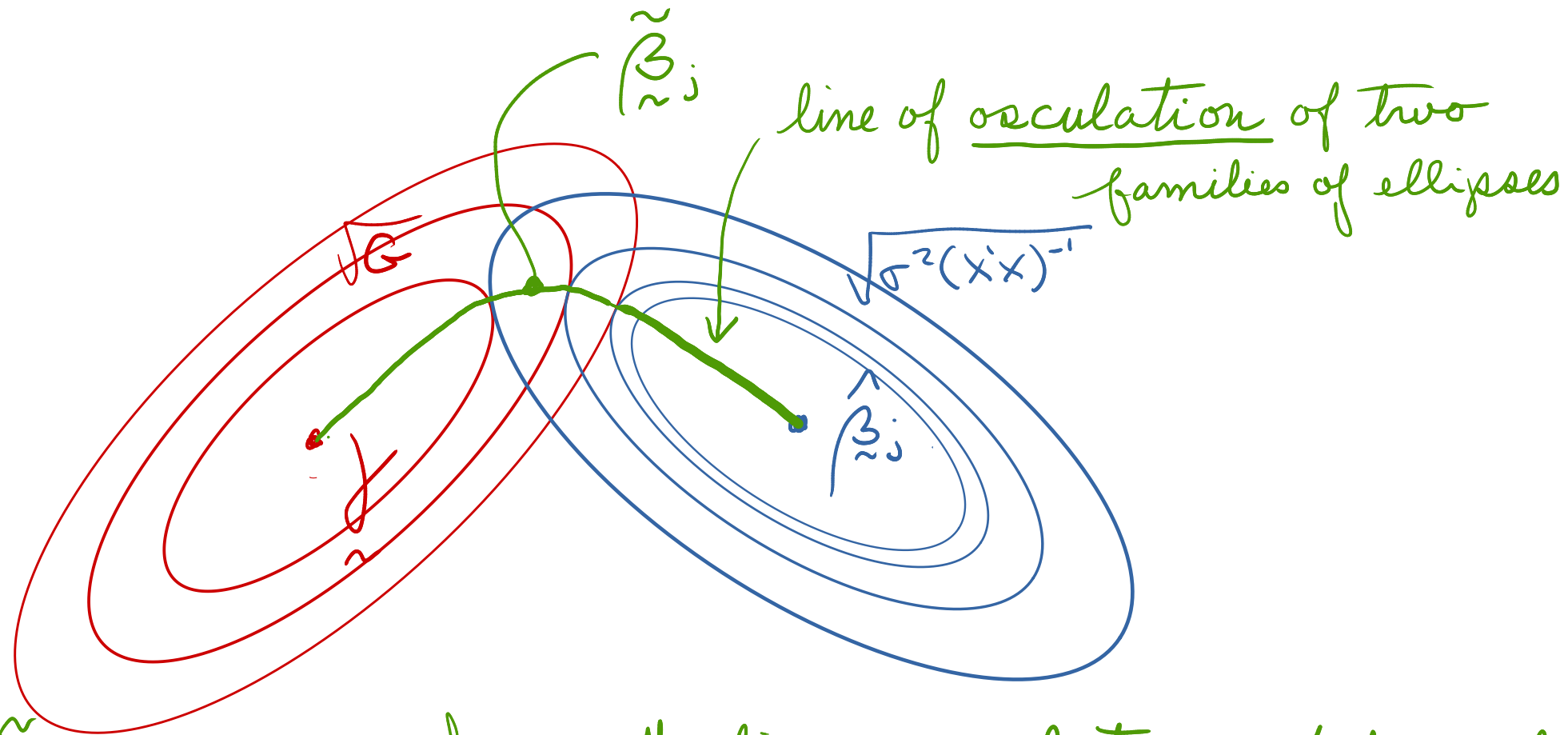
$$W_1 = G^{-1}$$

$$W_2 = [\sigma^2 (X_j' X_j)^{-1}]^{-1} \\ = \frac{1}{\sigma^2} X_j' X_j$$

What does this look like?

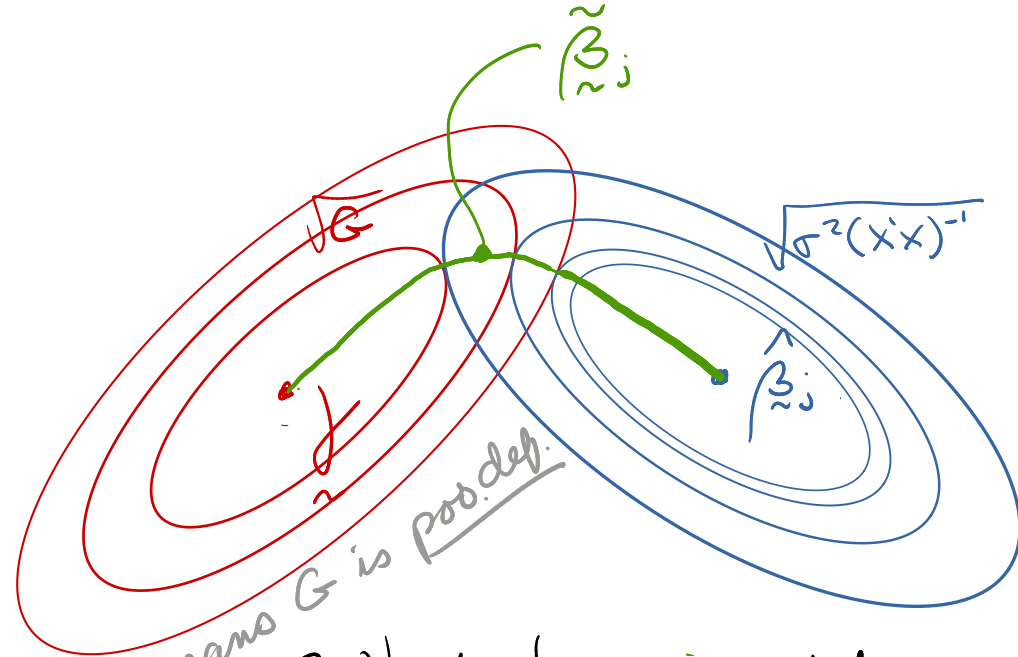






$\tilde{\beta}_{\tilde{j}}$  will be somewhere on the line of osculation of the families

$$\tilde{\gamma} \oplus k \sqrt{G} \quad \text{and} \quad \tilde{\beta} \oplus k \sqrt{\sigma^2(X'X)^{-1}}$$



- This is often referred to as "shrinking"  $\hat{\beta}_j$  towards  $\gamma$
- $\tilde{\beta}_j$  is a <sup>type of</sup> "shrinkage" estimator

What happens if :   
 asymptotic inner product matrix

$$\left. \begin{array}{l} n_j \rightarrow \infty \\ \text{and } G > 0 \end{array} \right\} X'X \rightarrow n\Pi_X \quad \text{so } \frac{1}{\sqrt{n}} \sqrt{\sigma^2 \Pi} \rightarrow 0 \quad \text{and } \tilde{\beta}_j \rightarrow \hat{\beta}_j$$

$G$  nearly singular with one small radius  $\} \tilde{\beta}_j$  will be close to the space  $\perp$  small radius.

$n_j$  is small  $\} \tilde{\beta}_j$  will be dominated by  $\gamma$





































































































Q

